



(11)

EP 3 072 129 B1

(12)

EUROPEAN PATENT SPECIFICATION

(45) Date of publication and mention of the grant of the patent:
13.06.2018 Bulletin 2018/24

(51) Int Cl.:
G10L 21/0208 ^(2013.01)

(21) Application number: **14721355.7**

(86) International application number:
PCT/EP2014/058913

(22) Date of filing: **30.04.2014**

(87) International publication number:
WO 2015/165539 (05.11.2015 Gazette 2015/44)

(54) **SIGNAL PROCESSING APPARATUS, METHOD AND COMPUTER PROGRAM FOR DEREVERBERATING A NUMBER OF INPUT AUDIO SIGNALS**

SIGNALVERARBEITUNG VORRICHTUNG, VERFAHREN UND COMPUTER PROGRAMM ZUR ENTHALLUNG EINER ANZAHL VON EINGANGSAUDIOSIGNALEN

DISPOSITIF, PROCEDE ET PROGRAMME INFORMATIQUE POUR LA DEREVERBERATION D'UN NOMBRE D'ENTREES DE SIGNAL AUDIO

(84) Designated Contracting States:
AL AT BE BG CH CY CZ DE DK EE ES FI FR GB GR HR HU IE IS IT LI LT LU LV MC MK MT NL NO PL PT RO RS SE SI SK SM TR

(43) Date of publication of application:
28.09.2016 Bulletin 2016/39

(73) Proprietor: **Huawei Technologies Co., Ltd.**
Longgang District
Shenzhen, Guangdong 518129 (CN)

(72) Inventors:
• **HELWANI, Karim**
80992 Munich (DE)
• **PANG, Liyun**
80992 Munich (DE)

(74) Representative: **Kreuz, Georg Maria**
Huawei Technologies Duesseldorf GmbH
Riesstrasse 8
80992 München (DE)

(56) References cited:
• **RASHOBH RAJAN S ET AL: "Multichannel Equalization in the KLT and Frequency Domains With Application to Speech Dereverberation", IEEE/ACM TRANSACTIONS ON AUDIO, SPEECH, AND LANGUAGE PROCESSING, IEEE, USA, vol. 22, no. 3, 1 March 2014 (2014-03-01), pages 634-646, XP011538165, ISSN: 2329-9290, DOI: 10.1109/TASLP.2013.2297013 [retrieved on 2014-01-24]**

- **ANDREAS WALTHER ET AL: "Direct-ambient decomposition and upmix of surround signals", APPLICATIONS OF SIGNAL PROCESSING TO AUDIO AND ACOUSTICS (WASPAA), 2011 IEEE WORKSHOP ON, IEEE, 16 October 2011 (2011-10-16), pages 277-280, XP032011488, DOI: 10.1109/ASPAA.2011.6082279 ISBN: 978-1-4577-0692-9 cited in the application**
- **HELWANI KARIM ET AL: "Multichannel acoustic echo suppression", 2013 IEEE INTERNATIONAL CONFERENCE ON ACOUSTICS, SPEECH AND SIGNAL PROCESSING (ICASSP); VANCOUCER, BC; 26-31 MAY 2013, INSTITUTE OF ELECTRICAL AND ELECTRONICS ENGINEERS, PISCATAWAY, NJ, US, 26 May 2013 (2013-05-26), pages 600-604, XP032508438, ISSN: 1520-6149, DOI: 10.1109/ICASSP.2013.6637718 [retrieved on 2013-10-18]**
- **Andreas Schwarz ET AL: "Coherence-based Dereverberation for Automatic Speech Recognition", 40th Annual German Congress on Acoustics DAGA 2014, 10 March 2014 (2014-03-10), pages 1-2, XP055153709, Oldenburg Retrieved from the Internet: URL: https://andreas-s.net/papers/schwarz_daga2014.pdf [retrieved on 2014-11-18]**

Note: Within nine months of the publication of the mention of the grant of the European patent in the European Patent Bulletin, any person may give notice to the European Patent Office of opposition to that patent, in accordance with the Implementing Regulations. Notice of opposition shall not be deemed to have been filed until the opposition fee has been paid. (Art. 99(1) European Patent Convention).

EP 3 072 129 B1

- **TAKUYA YOSHIOKA ET AL:** "Blind Separation and Dereverberation of Speech Mixtures by Joint Optimization", IEEE TRANSACTIONS ON AUDIO, SPEECH AND LANGUAGE PROCESSING, IEEE SERVICE CENTER, NEW YORK, NY, USA, vol. 19, no. 1, 1 January 2011 (2011-01-01), pages 69-84, XP011304608, ISSN: 1558-7916, DOI: 10.1109/TASL.2010.2045183
- **BENESTY J ET AL:** "A Blind Channel Identification-Based Two-Stage Approach to Separation and Dereverberation of Speech Signals in a Reverberant Environment", IEEE TRANSACTIONS ON SPEECH AND AUDIO PROCESSING, IEEE SERVICE CENTER, NEW YORK, NY, US, vol. 13, no. 5, 1 September 2005 (2005-09-01), pages 882-895, XP011137540, ISSN: 1063-6676, DOI: 10.1109/TSA.2005.851941
- **WANG LONGBIAO ET AL:** "Speech recognition using blind source separation and dereverberation method for mixed sound of speech and music", 2013 ASIA-PACIFIC SIGNAL AND INFORMATION PROCESSING ASSOCIATION ANNUAL SUMMIT AND CONFERENCE, APSIPA, 29 October 2013 (2013-10-29), pages 1-4, XP032549634, DOI: 10.1109/APSIPA.2013.6694159 [retrieved on 2013-12-24]

DescriptionTECHNICAL FIELD

5 **[0001]** The invention relates to the field of audio signal processing, in particular to the field of dereverberation and audio source separation.

BACKGROUND OF THE INVENTION

10 **[0002]** Dereverberation and audio source separation is a major challenge in a number of applications, such as multi-channel audio acquisition, speech acquisition, or up-mixing of mono-channel audio signals. Applicable techniques can be classified into single-channel techniques and multi-channel techniques.

[0003] Single-channel techniques can be based on a minimum statistics principle and can estimate an ambient part and a direct part of the audio signal separately. Single-channel techniques can further be based on a statistical system model. Common single-channel techniques, however, suffer from a limited performance in complex acoustic scenarios and may not be generalized to multi-channel scenarios.

15 **[0004]** Multi-channel techniques can aim at inverting a multiple input / multiple output finite impulse response (MIMO FIR) system between a number of audio signal sources and microphones, wherein each acoustic path between an audio signal source and a microphone can be modelled by an FIR filter. Multi-channel techniques can be based on higher order statistics and can employ heuristic statistical models using training data. Common multi-channel techniques, however, suffer from a high computational complexity and may not be applicable in single-channel scenarios.

20 **[0005]** In the document Herbert Buchner et al., "Trinicon for dereverberation of speech and audio signals", Speech Dereverberation, Signals and Communication Technology, pages 311-385, Springer London, 2010, an approach to estimate an ideal inverse system is described. In the document Andreas Walther et al., "Direct-Ambient Decomposition and Upmix of Surround Signals", IEEE Workshop on Applications of Signal Processing to Audio and Acoustics, 2011, an approach to estimate diffuse and direct audio component is described. In the document R.S. Rashobh, A.W.H. Khong, D. Liu "Multichannel Equalization in the KLT and Frequency Domains With Application to Speech Dereverberation", IEEE/ACM Transactions on Audio, Speech, and Language Processing, Vol.22, No.3, March 2014, three equalization algorithms based on a transform and a frequency domain are proposed.

30

SUMMARY OF THE INVENTION

[0006] It is an object of the invention to provide an efficient concept for dereverberating a number of input audio signals. The concept can also be applied for audio source separation within the number of input audio signals.

35 **[0007]** This object is achieved by the features of the independent claims. Further implementation forms are apparent from the dependent claims, the description and the figures.

[0008] Aspects and implementation forms of the invention are based on the finding that a filter coefficient matrix can be designed in a way that each output audio signal is coherent to its own history within a set of consequent time intervals and orthogonal to the history of other audio source signals. The filter coefficient matrix can be determined upon the basis of an initial guess of the audio source signals or upon the basis of a blind estimation approach. The invention can be applied using single-channel audio signals as well as multi-channel audio signals.

40 **[0009]** According to a first aspect, the invention relates to a signal processing apparatus for dereverberating a number of input audio signals according to claim 1. The number of input audio signals can be one or more than one. Thus, an efficient concept for dereverberation and/or audio source separation can be realized.

45 **[0010]** In an implementation form of the apparatus according to the first aspect as such, the filter coefficient determiner is configured to determine the signal space upon the basis of an input auto correlation matrix of the input transformed coefficient matrix. Thus, the signal space can be determined upon the basis of correlation characteristics of the input audio signals.

50 **[0011]** In an implementation form of the apparatus according to the first aspect, the transformer is configured to transform the number of input audio signals into frequency domain to obtain the input transformed coefficients. Thus, frequency domain characteristics of the input audio signals can be used to obtain the input transformed coefficients. The input transformed coefficients can relate to a frequency bin, e.g. having an index k, of a discrete Fourier transform (DFT) or a fast Fourier transform (FFT).

55 **[0012]** In an implementation form of the apparatus according to the first aspect, the transformer is configured to transform the number of input audio signals into the transformed domain for a number of past time intervals to obtain the input transformed coefficients. Thus, time domain characteristics of the input audio signals within a current time interval and past time intervals can be used to obtain the input transformed coefficients. The input transformed coefficients can relate to a time interval, e.g. having an index n, of a short time Fourier transform (STFT).

[0013] In an implementation form of the apparatus according to the first aspect, the filter coefficient determiner is configured to determine the filter coefficient matrix according to the following equation:

$$\mathbf{H} = \Phi_{xx}^{-1} \Gamma_{xS_0} \cdot \left(\Gamma_{xS_0}^H \Phi_{xx}^{-1} \Gamma_{xS_0} \right)^{-1}$$

wherein $\mathbf{H}$ denotes the filter coefficient matrix, $\mathbf{x}$ denotes the input transformed coefficient matrix, $\mathbf{S}_0$ denotes an auxiliary transformed coefficient matrix, Φ_{xx} denotes an input auto correlation matrix of the input transformed coefficient matrix, Γ_{xS_0} denotes a cross coherence matrix between the input transformed coefficient matrix and the auxiliary transformed coefficient matrix. Thus, the filter coefficient matrix can be determined efficiently upon the basis of an initial guess of the auxiliary transformed coefficient matrix.

[0014] In an implementation form of the apparatus according to the first aspect, the signal processing apparatus further comprises an auxiliary audio signal generator being configured to generate a number of auxiliary audio signals upon the basis of the number of input audio signals, and a further transformer being configured to transform the number of auxiliary audio signals into the transformed domain to obtain auxiliary transformed coefficients, the auxiliary transformed coefficients being arranged to form the auxiliary transformed coefficient matrix. Thus, the auxiliary transformed coefficient matrix can be determined upon the basis of the input audio signals.

[0015] The auxiliary audio signal generator can generate the number of auxiliary audio signals using a beamforming technique, e.g. a delay-and-sum beamforming technique, and/or by using audio signals of spot microphones. The auxiliary audio signal generator can therefore provide for an initial separation of a number of audio sources.

[0016] In an implementation form of the apparatus according to the first aspect, the filter coefficient determiner is configured to determine the filter coefficient matrix according to the following equation:

$$\mathbf{H} = \Phi_{xx}^{-1} \hat{\Gamma}_{sS} \cdot \left(\hat{\Gamma}_{sS}^H \Phi_{xx}^{-1} \hat{\Gamma}_{sS} \right)^{-1}$$

wherein $\mathbf{H}$ denotes the filter coefficient matrix, $\mathbf{x}$ denotes the input transformed coefficient matrix, Φ_{xx} denotes an input auto correlation matrix of the input transformed coefficient matrix, and $\hat{\Gamma}_{sS}$ denotes an estimate auto coherence matrix. Thus, the filter coefficient matrix can be determined efficiently upon the basis of an estimate auto coherence matrix.

[0017] In an implementation form of the apparatus according to the first aspect, the filter coefficient determiner is configured to determine the estimate auto coherence matrix according to the following equation:

$$\hat{\Gamma}_{sS}(k, n) := \left(\mathbf{I}_M \otimes \mathbf{U}^{-1} \right) \cdot \Gamma_{xx} \cdot \mathbf{U}$$

wherein $\hat{\Gamma}_{sS}$ denotes the estimate auto coherence matrix, $\mathbf{x}$ denotes the input transformed coefficient matrix, Γ_{xx} denotes an input auto coherence matrix of the input transformed coefficient matrix, $\mathbf{I}_M$ denotes an identity matrix of matrix dimension M , $\mathbf{U}$ denotes an eigenvector matrix of an eigenvalue decomposition performed upon the basis of the input auto coherence matrix. Thus, the estimate auto coherence matrix can efficiently be determined upon the basis of an eigenvalue decomposition.

[0018] In an implementation form of the apparatus according to the first aspect, the signal processing apparatus further comprises a channel determiner being configured to determine channel transformed coefficients upon the basis of the input transformed coefficients of the input transformed coefficient matrix and the filter coefficients of the filter coefficient matrix, the channel transformed coefficients being arranged to form a channel transformed matrix. Thus, a blind channel estimation can be performed.

[0019] In an implementation form of the apparatus according to the first aspect, the channel determiner is configured to determine the channel transformed matrix according to the following equation:

$$\hat{\mathbf{G}}(k, n) = \left(\mathbf{H}^H \mathbf{x}(k, n) \text{diag}\{X_1(k, n), X_2(k, n), \dots, X_P(k, n)\}^{-1} \right)^{-1}$$

wherein $\hat{\mathbf{G}}$ denotes the channel transformed matrix, $\mathbf{x}$ denotes the input transformed coefficient matrix, $\mathbf{H}$ denotes the filter coefficient matrix, and X_1 to X_P denote input transformed coefficients. Thus, the channel transformed matrix can be determined efficiently.

[0020] In an implementation form of the apparatus according to the first aspect, the number of input audio signals comprise audio signal portions being associated to a number of audio signal sources, and the signal processing apparatus is configured to separate the number of audio signal sources upon the basis of the number of input audio signals. Thus, a dereverberation and/or audio source separation can be performed.

[0021] According to a second aspect, the invention relates to a signal processing method for dereverberating a number of input audio signals according to claim 6. The number of input audio signals can be one or more than one. Thus, an efficient concept for dereverberation and/or audio source separation can be realized.

[0022] The signal processing method can be performed by the signal processing apparatus. Further features of the signal processing method can directly result from the functionality of the signal processing apparatus.

[0023] In an implementation form of the method according to the second aspect, the signal processing method further comprises determining the signal space upon the basis of an input auto correlation matrix of the input transformed coefficient matrix. Thus, the signal space can be determined upon the basis of correlation characteristics of the input audio signals.

[0024] According to a third aspect, the invention relates to a computer program comprising a program code for performing the signal processing method according to the second aspect as such or any implementation form of the second aspect when executed on a computer. Thus, the method can be performed in an automatic and repeatable manner.

[0025] The computer program can be provided in form of a machine-readable code. The computer program can comprise a series of commands for a processor of the computer. The processor of the computer can be configured to execute the computer program. The computer can comprise a processor, a memory, and/or input/output means.

[0026] The invention can be implemented in hardware and/or software.

[0027] Further embodiments of the invention will be described with respect to the following figures, in which:

Fig. 1 shows a diagram of a signal processing apparatus for dereverberating a number of input audio signals according to an implementation form;

Fig. 2 shows a diagram of a signal processing method for dereverberating a number of input audio signals according to an implementation form;

Fig. 3 shows a diagram of a signal processing apparatus for dereverberating a number of input audio signals according to an implementation form;

Fig. 4 shows a diagram of an audio signal acquisition scenario according to an implementation form;

Fig. 5 shows a diagram of a structure of an auto coherence matrix according to an implementation form;

Fig. 6 shows a diagram of a structure of an intermediate matrix according to an implementation form;

Fig. 7 shows a spectrogram of an input audio signal and a spectrogram of an output audio signal according to an implementation form; and

Fig. 8 shows a diagram of a signal processing apparatus for dereverberating a number of input audio signals according to an implementation form.

DETAILED DESCRIPTION OF EMBODIMENTS OF THE INVENTION

[0028] Fig. 1 shows a diagram of a signal processing apparatus 100 for dereverberating a number of input audio signals according to an implementation form.

[0029] The signal processing apparatus 100 comprises a transformer 101 being configured to transform the number of input audio signals into a transformed domain to obtain input transformed coefficients, the input transformed coefficients being arranged to form an input transformed coefficient matrix, a filter coefficient determiner 103 being configured to determine filter coefficients upon the basis of eigenvalues of a signal space, the filter coefficients being arranged to form a filter coefficient matrix, a filter 105 being configured to convolve input transformed coefficients of the input transformed coefficient matrix by filter coefficients of the filter coefficient matrix to obtain output transformed coefficients, the output transformed coefficients being arranged to form an output transformed coefficient matrix, and an inverse transformer 107 being configured to inversely transform the output transformed coefficient matrix from the transformed domain to obtain a number of output audio signals.

[0030] Fig. 2 shows a diagram of a signal processing method 200 for dereverberating a number of input audio signals according to an implementation form.

[0031] The signal processing method 200 comprises transforming 201 the number of input audio signals into a transformed domain to obtain input transformed coefficients, the input transformed coefficients being arranged to form an input transformed coefficient matrix, determining 203 filter coefficients upon the basis of eigenvalues of a signal space, the filter coefficients being arranged to form a filter coefficient matrix, convolving 205 input transformed coefficients of

the input transformed coefficient matrix by filter coefficients of the filter coefficient matrix to obtain output transformed coefficients, the output transformed coefficients being arranged to form an output transformed coefficient matrix, and inversely transforming 207 the output transformed coefficient matrix from the transformed domain to obtain a number of output audio signals.

[0032] The signal processing method 200 can be performed by the signal processing apparatus 100. Further features of the signal processing method 200 can directly result from the functionality of the signal processing apparatus 100 as described above and below in further detail.

[0033] Fig. 3 shows a diagram of a signal processing apparatus 100 for dereverberating a number of input audio signals according to an implementation form. The signal processing apparatus 100 comprises a transformer 101, a filter coefficient determiner 103, a filter 105, an inverse transformer 107, an auxiliary audio signal generator 301, a further transformer 303, and a post-processor 305.

[0034] The transformer 101 can be a short time Fourier transform (STFT) transformer. The filter coefficient determiner 103 can perform an algorithm. The filter 105 can be characterized by a filter coefficient matrix H . The inverse transformer 107 can be an inverse short time Fourier transform (ISTFT) transformer. The auxiliary audio signal generator 301 can provide an initial guess, e.g. by using a delay-and-sum technique and/or spot microphone audio signals. The further transformer 303 can be a short time Fourier transform (STFT) transformer. The post-processor 305 can provide post-processing capabilities, e.g. an automatic speech recognition (ASR), and/or an up-mixing.

[0035] A number Q of input audio signals can be provided to the transformer 101 and the auxiliary audio signal generator 301. The auxiliary audio signal generator 301 can provide a number of P auxiliary audio signals to the further transformer 303. The further transformer 303 can provide a number P of rows or columns of an auxiliary transformed coefficient matrix to the filter coefficient determiner 103. The filter 105 can provide a number P of rows or columns of an output transformed coefficient matrix to the inverse transformer 107. The inverse transformer 107 can provide a number P of output audio signals to the post-processor 305 yielding a number P of post-processed audio signals.

[0036] The diagram shows an overall architecture of the apparatus 100. The input to the apparatus 100 can be microphone signals. These can optionally be preprocessed by an algorithm offering spatial selectivity, e.g. a delay-and-sum beamformer. The preprocessed signals and/or microphone signals can be analyzed by an STFT. The microphone signals can then be stored in a buffer with optionally variable size for the different frequency bins. The algorithms can calculate filter coefficients based on the buffered audio signal time intervals or frames. The buffered signal can be filtered in each frequency bin with a calculated complex filter. The output of the filtering can be transformed back to the time domain. The processed audio signals can optionally be fed into the post-processor 305, such as for automatic speech recognition (ASR) or up-mixing.

[0037] Some implementation forms can relate to blind single-channel and/or multi-channel minimization of an acoustical influence of an unknown room. They can be employed in multi-channel acquisition systems in telepresence for enhancing the ability of the systems to focus onto a part of a captured acoustic scene, speech and signal enhancement for mobiles and tablets, in particular by dereverberation of signals in a hands-free mode, and also for up-mixing of mono signals.

[0038] For this purpose, an approach for blind dereverberation and/or source separation can be used. The approach can be specialized to a single-channel case and can be used as a blind source separation post-processing stage.

[0039] The propagation of sound waves from a sound source to a predefined measurement point under typical conditions can be described by convolving the sound source signal with a Green's function which can solve an inhomogeneous wave equation under given boundary conditions. The boundary conditions, however, may not be controllable and may result in undesired acoustic characteristics such as long reverberation time which can cause insufficient intelligibility. In advanced communication systems which are able to synthesize a user defined acoustic environment, it can be desirable to mitigate the influence of the recording room and to maintain only a clean excitation signal to integrate it properly in the desired virtual acoustic environment.

[0040] In the case of multiple sound sources, e.g. speakers, captured by a distributed microphone array in a recording room, dereverberation can offer original clean source signals separated and free of the recording room influence, e.g. speech signals as would be recorded by a microphone next to the mouth of a single speaker in an anechoic chamber.

[0041] Dereverberation techniques can aim at minimizing the effect of the late part of the room impulse response. However, a full deconvolution of the microphone signals can be challenging and the output can be a less reverberant mixture of the source signals but not separated source signals.

[0042] Dereverberation techniques can be classified into single-channel and multi-channel techniques. Due to theoretical limits, an ideal deconvolution can typically be achieved in the multi-channel case where the number of recording microphones Q can be higher than the number of active sound sources P , e.g. speakers.

[0043] Multi-channel dereverberation techniques can aim at inverting a multiple input/output, finite impulse response, i.e. MIMO FIR, system between the sound sources and the microphones wherein each acoustic path between a sound source and a microphone can be modelled by an FIR filter of length L . The MIMO system can be presented in time domain as a matrix that can be invertible if it is square and regular. Hence, an ideal inversion can be performed if the following two conditions hold.

[0044] Firstly, the length L' of a finite inverse filter fulfils:

$$L' = \frac{P(L-1)}{Q-P} \quad (1)$$

[0045] Secondly, the individual filters of the MIMO system do not exhibit common roots in the z-domain.

[0046] An approach to estimate an ideal inverse system can be employed. The approach can be based on exploiting a non-Gaussianity, a non-whiteness, and a non-stationarity of the source signals. The approach can feature a minimum distortion on the cost of a high computational complexity for the computation of higher order statistics. Moreover, since it can aim at solving an ideal inversion problem, it may require from the system to have more microphones than sound sources and may not be applicable for a single channel problem.

[0047] A further approach to dereverberate a multi-channel recording can be based on estimating a signal subspace. Ambient and direct parts of the audio signal can be estimated separately. Late reverberations can be estimated and can be treated as noise. Therefore, the approach may require an accurate estimation of the ambient part, i.e. the late reverberations, to be able to cancel it. The approaches based on estimating a multi-channel signal subspace can be dedicated to reduce the reverberance and not to de-mix, i.e. to separate, the sound sources. The approaches are typically applied to multi-channel setups and may not be used to solve a single channel dereverberation problem. Additionally, heuristic statistical models to estimate the reverberation and to reduce the ambient part can be employed. These models may be based on training data and may suffer from a high complexity.

[0048] A further approach to estimate diffuse and direct components in the spectral domain can be employed. The short-time spectra of a multi-channel signal can be down-mixed into $X_1(k,n)$ and $X_2(k,n)$, wherein k and n denote a frequency bin index and a time interval or frame index. A real coefficient $H(k,n)$ can be derived to extract the direct components $\hat{S}_1(k,n)$ and $\hat{S}_2(k,n)$ from the down-mix according to:

$$\hat{S}_1(k,n) = H(k,n) \cdot X_1(k,n)$$

$$\hat{S}_2(k,n) = H(k,n) \cdot X_2(k,n)$$

[0049] Under the assumption that direct and diffuse components in the down-mix are mutually uncorrelated and the diffuse components in the down-mix have equal power, the real coefficient $H(k,n)$ can be calculated based on a Wiener optimization criterion according to:

$$H(k,n) = \frac{P_S}{P_S + P_A}$$

wherein P_S and P_A are the sums of the short-time power spectral estimates of the direct and diffuse components in the

down-mix. P_S and P_A can be derived based on the cross-correlation of the down-mix as $Re(E\{X_1 X_2^*\})$. These filters can further be applied to multi-channel audio signals to generate the corresponding direct and ambient components. This approach can be based on a multi-channel setup and may not solve a single channel dereverberation problem. Moreover, it may introduce a high amount of distortion and may not perform a de-mixing.

[0050] Single channel dereverberation solutions can be based on the minimum statistics principle. Therefore, they may estimate the ambient and the direct part of the audio signal separately. An approach that incorporates a statistical system model can be employed which can be based on training data. A further approach can be applied on a single channel setup offering limited performance in complex sound scenes, especially with respect to the audio signal quality since the approach can be optimized for automatic speech recognition and not for a high quality listening experience.

[0051] Some implementation forms can relate to single-channel and multi-channel dereverberation techniques. In order to obtain a dry output audio signal, an M-taps MIMO FIR filter in the STFT domain with P outputs, i.e. number of audio signal sources, and Q inputs, i.e. number of input audio signals, number of microphones, or number of outputs of a preprocessing stage such as a beamformer, e.g. a delay-and-sum beamformer, can be applied. The filter can be designed in a way that each output audio signal can be coherent to its own history within a predefined set of consequent time intervals or frames and can be orthogonal to the history of the other audio source signals.

[0052] In the following, a mathematical setup and a signal model is introduced used to derive the dereverberation approach. The input audio signal x_q at a time instant t can be given as a convolution of a dry excitation audio source signal $s(t) := [s_1(t), s_2(t), \dots, s_P(t)]^T$ convolved with Green's functions for the p^{th} source to the q^{th} input or microphone $g_q(t) := [g_{1q}, g_{2q}, \dots, g_{Pq}]^T$:

$$x_q(t) = \sum_{p=1}^P s_p(t) * g_{pq}(x) \quad (2)$$

[0053] By considering this equation in the short time Fourier domain, it can be approximated as:

$$X_q(k, n) \approx [S_1, S_2, \dots, S_P] \cdot [G_{1q}, G_{2q}, \dots, G_{Pq}]^H, \quad (3)$$

wherein k denotes a frequency bin index and the time interval or frame is indexed by n , $\{\cdot\}^H$ denotes a Hermitian transpose, and the dependencies of both the audio signal source signals and the Green's functions on (n, k) are avoided for clarity of notation. For a complete multi-channel representation, it can be written for the MIMO system:

$$\mathbf{X}(k, n) \approx [S_1, S_2, \dots, S_P] \cdot \begin{bmatrix} G_{11} & \dots & G_{P1} \\ \vdots & \ddots & \vdots \\ G_{1Q} & \dots & G_{PQ} \end{bmatrix}^H, \quad (4)$$

$$\mathbf{X}(k, n) \approx \mathbf{S}^T(k, n) \cdot \mathbf{G}^H(k, n),$$

with

$$\mathbf{X} := [X_1(k, n), X_2(k, n), \dots, X_Q(k, n)]^T, \quad (5)$$

$$\mathbf{S} := [S_1(k, n), S_2(k, n), \dots, S_P(k, n)]^T, \quad (6)$$

$$\mathbf{G} := \begin{bmatrix} G_{11} & \dots & G_{P1} \\ \vdots & \ddots & \vdots \\ G_{1Q} & \dots & G_{PQ} \end{bmatrix}. \quad (7)$$

[0054] A dereverberation can be performed using an FIR filter in the STFT domain, for example based on applying an FIR filter according to:

$$\mathbf{H}(k, n) := \begin{bmatrix} \mathbf{h}_{11}(k, n) & \dots & \dots & \dots & \mathbf{h}_{P1}(k, n) \\ \vdots & \ddots & \ddots & \ddots & \vdots \\ \vdots & \mathbf{h}_{pq}(k, n) & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & \ddots & \vdots \\ \mathbf{h}_{1Q}(k, n) & \dots & \dots & \dots & \mathbf{h}_{PQ}(k, n) \end{bmatrix}, \quad (8)$$

with $h_{pq}(k, n) := [H_{pq}(k, n), H_{pq}(k, n-1), \dots, H_{pq}(k, n-M+1)]^T$ in the STFT domain on the input audio signal

$$\hat{\mathbf{S}}(k, n) := \mathbf{H}^H(k, n) \mathbf{x}(k, n), \quad (9)$$

wherein a sequence of M consecutive STFT domain time intervals or frames of the input audio signal is defined as:

$$\mathbf{x}_q(k, n) := [X_q(k, n), X_q(k, n-1), \dots, X_q(k, n-M+1)]^T, \quad (10)$$

and

$$\mathbf{x}(k, n) := [\mathbf{x}_1^T(k, n), \mathbf{x}_2^T(k, n), \dots, \mathbf{x}_q^T(k, n), \dots, \mathbf{x}_Q^T(k, n)]^T, \quad (11)$$

$$\hat{\mathbf{S}}(k, n) := [\hat{S}_1(k, n), \hat{S}_2(k, n), \dots, \hat{S}_P(k, n)]^T. \quad (12)$$

[0055] Note that M can be chosen individually for each frequency bin. For example, for a speech signal using a sampling frequency of 16 kHz, a STFT window size of 320, a STFT length of 512, an overlapping factor of 0.5, and a reverberation time of approximately 1 second, M can be set to 4 for the lower 129 bins, and can be set to 2 for the higher 128 bins.

[0056] The filter coefficient matrix H can approximate the largest eigenvectors of the auto correlation matrix of the unknown dry audio source signal. It can be desirable to obtain a distortionless estimate of the dry audio source signal. This can mean that the FIR filter exhibits fidelity to the coherent part of the dry audio source signal.

[0057] The input audio signal can be decomposed into a part which is coherent with an initial estimation of the dry audio source signal $\mathbf{x}_c$, and an incoherent part $\mathbf{x}_i$ according to:

$$\mathbf{x}(k, n) = \mathbf{x}_c(k, n) + \mathbf{x}_i(k, n), \quad (13)$$

with

$$\mathbf{x}_c(k, n) := \mathbf{\Gamma}_{\mathbf{xS}}(k, n) \cdot \mathbf{S}(k, n), \quad (14)$$

wherein a cross coherence matrix of the dry audio source signal can be defined as a normalized correlation matrix by:

$$\mathbf{\Gamma}_{\mathbf{xS}}(k, n) := \hat{\mathcal{E}}\{\mathbf{x}(k, n) \mathbf{S}^H(k, n)\} \cdot (\phi_{\mathbf{SS}}(k, n))^{-1}, \quad (15)$$

wherein $\hat{\mathcal{E}}\{\cdot\}$ denotes an estimation of an expectation value, and with the estimation of the expectation of auto correlation matrix

$$\phi_{\mathbf{SS}}(k, n) := \hat{\mathcal{E}}\{\mathbf{S}(k, n) \mathbf{S}^H(k, n)\}. \quad (16)$$

[0058] The cross coherence matrix $\mathbf{\Gamma}_{\mathbf{xS}}$ can be understood as enforced eigenvectors matrix of the auto correlation matrix of the input audio signal.

[0059] The estimation of the expectation value can be calculated iteratively by

$$\hat{\mathcal{E}}\{\mathbf{x}(k, n) \mathbf{S}^H(k, n)\} = \alpha \hat{\mathcal{E}}\{\mathbf{x}(k, n-1) \mathbf{S}^H(k, n-1)\} + (1-\alpha) \mathbf{x}(k, n) \mathbf{S}^H(k, n) \quad (17)$$

$$\hat{\mathcal{E}}\{\mathbf{S}(k, n) \mathbf{S}^H(k, n)\} = \alpha \hat{\mathcal{E}}\{\mathbf{S}(k, n-1) \mathbf{S}^H(k, n-1)\} + (1-\alpha) \mathbf{S}(k, n) \mathbf{S}^H(k, n) \quad (18)$$

wherein α denotes a forgetting factor.

[0060] Hence, a condition for the dereverberation filter can be set as:

$$\mathbf{H}^H \hat{\mathcal{E}} \{ \mathbf{x}(k, n) \mathbf{S}^H(k, n) \} = \phi_{SS} \quad (19)$$

[0061] By rearranging, the following expression can be obtained:

$$\mathbf{H}^H \Gamma_{xS} = \mathbf{I}_{P \times P}, \quad (20)$$

wherein $\mathbf{I}$ denotes a unity matrix. Therefore, the filter coefficient matrix $\mathbf{H}$ can be coincident to the basis vectors Γ_{xS} of the signal subspace.

[0062] An optimal dereverberation FIR filter in the STFT domain can be derived. To obtain an optimal filter, the following cost function which can be constrained by (20) can be set:

$$J = \mathbf{H}^H \Phi_{xx} \mathbf{H} + \lambda \left(\mathbf{H}^H \Gamma_{xS} - \mathbf{I}_{P \times P} \right), \quad (21)$$

wherein

$$\Phi_{xx} := \hat{\mathcal{E}} \{ \mathbf{x} \mathbf{x}^H \} \quad (22)$$

wherein λ denotes a Lagrange multipliers matrix. At a minimum of this cost function, the gradient can be zero, and the optimal expression of the filter can be obtained as:

$$\mathbf{H} = \Phi_{xx}^{-1} \Gamma_{xS} \cdot \left(\Gamma_{xS}^H \Phi_{xx}^{-1} \Gamma_{xS} \right)^{-1}. \quad (23)$$

[0063] The filter can maximize the entropy of the dry audio signal under the given condition.

[0064] The cross coherence matrix can be approximated. In the following, two possibilities to deal with the missing unknown dry audio source signal are proposed.

[0065] Fig. 4 shows a diagram of an audio signal acquisition scenario 400 according to an implementation form. The audio signal acquisition scenario 400 comprises a first audio signal source 401, a second audio signal source 403, a third audio signal source 405, a microphone array 407, a first beam 409, a second beam 411, and a spot microphone 413. The first beam 409 and the second beam 411 are synthesized by the microphone array 407 by a beamforming technique.

[0066] The diagram shows the audio signal acquisition scenario 400 with three audio signal sources 401, 403, 405 or speakers, a microphone array 407 with the ability of achieving high sensitivity in dedicated directions, e.g. using beamforming, e.g. a delay-and-sum beamformer, and a spot microphone 413 next to one audio signal source. Separated audio sources 401, 403, 405 with a minimized room influence can be desired. The output of the beamformer and the auxiliary audio signal of the spot microphone 413 can be used to calculate or estimate the cross coherence matrix Γ_{xS} .

[0067] The algorithm can handle the output of the beamformer and of the spot microphone, i.e. the auxiliary audio signals, as an initial guess, enhance the separation and minimize the reverberation of the input audio signal or microphone array signal to provide a clean version of the three audio source signals or speech signals.

[0068] For calculating the derived filter coefficient matrix, a computation of a cross coherence matrix can be performed. Therefore, a pre-processing stage can be employed, e.g. a source localization stage combined with beamforming, providing an initial guess of the dry audio source signals $s_{01}, s_{02}, \dots, s_{0P}$, or even a combination with a spot microphone for a subset of the audio sources.

[0069] For the filter, the following expression can be obtained:

$$\mathbf{H} = \Phi_{xx}^{-1} \Gamma_{xS_0} \cdot \left(\Gamma_{xS_0}^H \Phi_{xx}^{-1} \Gamma_{xS_0} \right)^{-1}, \quad (24)$$

wherein Γ_{xS_0} can be defined by the same expression as in Eq. (15) but by using the initial guess instead of the dry audio

source signal.

[0070] Fig. 5 shows a diagram of a structure of an auto coherence matrix 501 according to an implementation form. The diagram shows a block-diagonal structure. The auto coherence matrix 501 can relate to Γ_{ss} . The auto coherence matrix 501 can comprise M x P rows and P columns.

[0071] Fig. 6 shows a diagram of a structure of an intermediate matrix 601 according to an implementation form. The diagram shows further an auto coherence matrix 603. The intermediate matrix 601 can relate to C. The intermediate matrix 601 or matrix C can be constructed based on a system with P=3 input audio signals or microphones. The auto coherence matrix 603 can comprise portions having M rows and can comprise Q columns. The auto coherence matrix 603 can relate to Γ_{xx} .

[0072] In the case $P = Q$, the condition in (20) can be modified for coherence of the output audio signals according to:

$$\mathbf{H}^H \Gamma_{ss} = \mathbf{I}_{P \times P} \quad (25)$$

[0073] For the case $P=Q$, it can be assumed that each source of the dry audio source signal is coherent with regard to its own history. Based on the assumptions, Γ_{ss} can be used instead of Γ_{xs} . Reverberations and interfering signals can be incoherent.

[0074] The auto coherence matrix of the audio source signal can be defined as:

$$\Gamma_{ss}(k, n) := \hat{\mathcal{E}}\{\mathbf{s}(k, n) \mathbf{S}^H(k, n)\} \cdot (\phi_{ss}(k, n))^{-1}, \quad (26)$$

wherein the quantity ϕ_{ss} can have a similar definition as (16):

$$\phi_{ss}(k, n) := \hat{\mathcal{E}}\{\mathbf{S}(k, n) \mathbf{S}^H(k, n)\}. \quad (27)$$

[0075] The auto coherence matrix Γ_{ss} of the audio sources can be block diagonal. Furthermore, in the spirit of Γ_{xs} an auto coherence matrix of the input audio signal can be introduced as:

$$\Gamma_{xx}(k, n) := \hat{\mathcal{E}}\{\mathbf{x}(k, n) \mathbf{X}^H(k, n)\} \cdot (\phi_{xx}(k, n))^{-1}, \quad (28)$$

wherein the quantity ϕ_{xx} can have a similar definition as (16):

$$\phi_{xx}(k, n) := \hat{\mathcal{E}}\{\mathbf{X}(k, n) \mathbf{X}^H(k, n)\}. \quad (29)$$

[0076] By assuming the Green's functions in (4) to be constant for the considered M time intervals or frames, it can be seen that:

$$\Gamma_{xx}(k, n) = \hat{\mathcal{E}}\{\mathbf{x}(k, n) \mathbf{S}^H(k, n)\} \cdot (\phi_{sx}(k, n))^{-1}, \quad (30)$$

with

$$\phi_{sx} := \hat{\mathcal{E}}\{\mathbf{S}(k, n) \mathbf{X}^H(k, n)\}. \quad (31)$$

[0077] In order to obtain an expression for Γ_{ss} , approximations can be made by assuming the audio source signals to be independent, i.e. ϕ_{ss} can be diagonal and $\hat{\mathcal{E}}\{\mathbf{s}(k, n) \mathbf{S}^H(k, n)\}$ can be block diagonal, and by taking into account the relation (30) for $P = Q$:

$$\Gamma_{xx}(k, n) = \mathbf{I}_M \otimes \mathbf{G}^* \cdot \hat{\mathcal{E}}\{\mathbf{s}(k, n) \mathbf{S}^H(k, n)\} \cdot (\phi_{sx}(k, n))^{-1}, \quad (32)$$

wherein $\otimes$ denotes a Kronecker product. Hence, in order to approximate Γ_{ss} , we can use Γ_{xx} and can set the off diagonal blocks to zero. This can be achieved by setting a square, non-necessarily symmetric, intermediate matrix C whose rows

are the $(j \cdot M + 1)^{th}$ row of the auto coherence matrix of the input audio signal, with $j \in \{0, \dots, P - 1\}$. Note, that the order may be maintained.

[0078] An eigenvalue decomposition can allow to write $\mathbf{C}$ as a product $\mathbf{U} \cdot \underline{\mathbf{C}} \cdot \mathbf{U}^{-1}$, wherein $\underline{\mathbf{C}}$ can be diagonal. An estimate $\hat{\Gamma}_{\mathbf{s}\mathbf{s}}(k, n)$ for the block diagonal form for Γ can be obtained as:

$$\hat{\Gamma}_{\mathbf{s}\mathbf{s}}(k, n) := (\mathbf{I}_M \otimes \mathbf{U}^{-1}) \cdot \Gamma_{\mathbf{x}\mathbf{x}} \cdot \mathbf{U} \quad (33)$$

[0079] To obtain a filter coefficient matrix that provides the coherent part of the audio signal sources, the following can be set similarly to Eq. (24):

$$\mathbf{H} = \Phi_{\mathbf{x}\mathbf{x}}^{-1} \hat{\Gamma}_{\mathbf{s}\mathbf{s}} \cdot (\hat{\Gamma}_{\mathbf{s}\mathbf{s}}^H \Phi_{\mathbf{x}\mathbf{x}}^{-1} \hat{\Gamma}_{\mathbf{s}\mathbf{s}})^{-1} \quad (34)$$

[0080] In addition, a blind channel estimation can be performed. An expression of the estimated inverse channel can be obtained by the following considerations for $X_P(k, n) \neq 0$:

$$\hat{\mathbf{S}}(k, n) = \mathbf{H}^H \mathbf{x}(k, n) \text{diag}\{X_1(k, n), X_2(k, n), \dots, X_P(k, n)\}^{-1} \cdot \text{diag}\{X_1(k, n), X_2(k, n), \dots, X_P(k, n)\}, \quad (35)$$

wherein the operator $\text{diag}\{\cdot\}$ creates a diagonal square matrix with an argument vector on the main diagonal. Comparing this equation to the assumed channel model in the STFT domain in (3) leads to:

$$\hat{\mathbf{G}}(k, n) = (\mathbf{H}^H \mathbf{x}(k, n) \text{diag}\{X_1(k, n), X_2(k, n), \dots, X_P(k, n)\}^{-1})^{-1} \quad (36)$$

[0081] Fig. 7 shows a spectrogram 701 of an input audio signal and a spectrogram 703 of an output audio signal according to an implementation form. In the spectrograms 701, 703, a magnitude of a corresponding short time Fourier transform (STFT) is color-coded over time in seconds and frequency in Hertz.

[0082] The spectrogram 701 can further relate to a reverberant microphone signal and the spectrogram 703 can further relate to an estimated dry audio source signal. In this example for a single channel, the spectrogram 701 of the reverberant signal is smeared out. Comparatively, the spectrogram 703 of the estimated dry audio source signal by applying the dereverberation algorithm exhibits a structure of a typical dry speech signal.

[0083] Fig. 8 shows a diagram of a signal processing apparatus 100 for dereverberating a number of input audio signals according to an implementation form. The signal processing apparatus 100 comprises a transformer 101, a filter coefficient determiner 103, a filter 105, an inverse transformer 107, an auxiliary audio signal generator 301, and a post-processor 305.

[0084] The transformer 101 can be a short time Fourier transform (STFT) transformer. The filter coefficient determiner 103 can perform an algorithm. The filter 105 can be characterized by a filter coefficient matrix $\mathbf{H}$. The inverse transformer 107 can be an inverse short time Fourier transform (ISTFT) transformer. The auxiliary audio signal generator 301 can provide an initial guess, e.g. by using a delay-and-sum technique and/or spot microphone audio signals. The post-processor 305 can provide post-processing capabilities, e.g. an automatic speech recognition (ASR), and/or an up-mixing.

[0085] A number Q of input audio signals can be provided to the auxiliary audio signal generator 301. The auxiliary audio signal generator 301 can provide a number P of auxiliary audio signals to the transformer 101. The transformer 101 can provide a number P of rows or columns of an input transformed coefficient matrix to the filter coefficient determiner 103 and the filter 105. The filter 105 can provide a number P of rows or columns of an output transformed coefficient matrix to the inverse transformer 107. The inverse transformer 107 can provide a number P of output audio signals to the post-processor 305 yielding a number P of post-processed audio signals.

[0086] The invention has several advantages. It can be used for post-processing for audio source separation achieving an optimal separation even with a low complexity solution for an initial guess. This can be used for enhanced sound-field recordings. It can further be used even for a single-channel dereverberation which can be a benefit to speech intelligibility for hands-free application using mobiles and tablets. It can further be used for up-mixing for multi-channel reproduction even from a mono recording and for pre-processing for automatic speech recognition (ASR).

[0087] Some implementation forms can relate to a method to modify a multi- or single-channel audio signal obtained by recording one or multiple audio signal sources in a reverberant acoustic environment, the method comprising mini-

mizing the influence of the reverberations caused by the room and separating the recorded audio sound sources. The recording can be done by a combination of a microphone array with the ability to perform pre-processing as localization of the audio signal sources and beamforming, e.g. delay-and-sum, and distributed microphones, e.g. spot microphones, next to a subgroup of the audio signal sources.

[0088] The non-preprocessed input audio signals or array signals and the pre-processed signals together with available distributed spot microphones can be analyzed using a short time Fourier transformation (STFT) and can be buffered. The length of the buffer, e.g. length M, can be chosen individually for each frequency band. The buffered input audio signals can be combined in the short time Fourier transformation domain to obtain 2-multidimensional complex filters for each sub-band that can exploit the inter time interval or inter-frame statistics of the audio signals. The dry output audio signals, i.e. the separated and/or dereverberated input audio signals, can be obtained by performing a multi-dimensional convolution of the input audio signals or array microphone signals with those filters. The convolution can be performed in the short time Fourier transformation domain.

[0089] The filters can be designed to fulfill the condition of maximum entropy of the output audio signals in the STFT domain constrained by maintaining the coherence, e.g. normalized cross correlation, between the pre-processed audio signal and the distributed spot microphones on one side and the input audio signals or array microphone signals on the other side according to:

$$\mathbf{H} = \Phi_{xx}^{-1} \Gamma_{xs0} \cdot \left(\Gamma_{xs0}^H \Phi_{xx}^{-1} \Gamma_{xs0} \right)^{-1}$$

[0090] Some implementation forms can further relate to a method wherein a pre-processing stage can be unavailable and the filters can be designed to maintain the coherence of each audio source signal to its own history and the independence of the audio signal sources in the STFT domain according to:

$$\mathbf{H} = \Phi_{xx}^{-1} \hat{\Gamma}_{sS} \cdot \left(\hat{\Gamma}_{sS}^H \Phi_{xx}^{-1} \hat{\Gamma}_{sS} \right)^{-1}$$

[0091] An estimate of an auto coherence matrix of the audio source signals can be calculated by means of an eigenvalue decomposition of a square matrix whose rows can be selected from the rows of an auto coherence of the input audio signals or microphone signals. The number of rows can be determined by the number of separable audio signal sources which may maximally be the number of inputs or microphones. The matrix U containing in its columns the eigenvectors of the so-constructed matrix C can be inverted and the estimate of the audio source auto coherence matrix can be calculated by:

$$\hat{\Gamma}_{sS}(k, n) := \left(\mathbf{I}_M \otimes \mathbf{U}^{-1} \right) \cdot \Gamma_{xx} \cdot \mathbf{U}$$

[0092] Some implementation forms can further relate to a method to estimate acoustic transfer functions based on the calculated optimal 2-dimensional filters according to:

$$\hat{\mathbf{G}}(k, n) = \left(\mathbf{H}^H \mathbf{x}(k, n) \text{diag}\{X_1(k, n), X_2(k, n), \dots, X_P(k, n)\}^{-1} \right)^{-1}$$

[0093] Some implementation forms can allow for a processing in the STFT domain. It can provide high system tracking capabilities because of an inherent batch block processing and high scalability, i.e. the resolution in time and frequency domain can freely be chosen by using suitable windows. The system can approximately be decoupled in the STFT domain. Therefore, the processing can be parallelized for each frequency bin. Furthermore, different sub-bands can be treated independently, e.g. different filter orders for dereverberation for different sub-bands can be used.

[0094] Some implementation forms can use a multi-tap approach in the STFT domain. Therefore, inter time interval or inter-frame statistics of the dry audio signals can be exploited. Each dry audio signal can be coherent to its own history. Therefore, it can be statistically represented over a predefined time by only one eigenvector. The eigenvectors of the audio source signals can be orthogonal.

Claims

1. A signal processing apparatus (100) for dereverberating a number (Q) of input audio signals, the signal processing

apparatus (100) comprising:

a transformer (101) being configured to transform the number (Q) of input audio signals into a transformed domain to obtain input transformed coefficients, the input transformed coefficients being arranged to form an input transformed coefficient matrix;
 a filter coefficient determiner (103) being configured to determine filter coefficients upon the basis of eigenvalues of a signal space, the filter coefficients being arranged to form a filter coefficient matrix (H);
 a filter (105) being configured to convolve input transformed coefficients of the input transformed coefficient matrix by filter coefficients of the filter coefficient matrix (H) to obtain output transformed coefficients, the output transformed coefficients being arranged to form an output transformed coefficient matrix; and
 an inverse transformer (107) being configured to inversely transform the output transformed coefficient matrix from the transformed domain to obtain a number of output audio signals; **characterised in that:** the filter coefficient determiner (103) is configured to determine input auto coherence coefficients upon the basis of the input transformed coefficients, the input auto coherence coefficients indicating a coherence of the input transformed coefficients associated to a current time interval and a past time interval, the input auto coherence coefficients being arranged to form an input auto coherence matrix, and wherein the filter coefficient determiner (103) is further configured to determine the filter coefficients upon the basis of the input auto coherence matrix.

2. The signal processing apparatus (100) of claim 1, wherein the filter coefficient determiner (103) is configured to determine the signal space upon the basis of an input auto correlation matrix (Φ_{xx}) of the input transformed coefficient matrix.
3. The signal processing apparatus (100) of any of the preceding claims, wherein the transformer (101) is configured to transform the number (Q) of input audio signals into frequency domain to obtain the input transformed coefficients.
4. The signal processing apparatus (100) of any of the preceding claims, further comprising:
 a channel determiner being configured to determine channel transformed coefficients upon the basis of the input transformed coefficients of the input transformed coefficient matrix and the filter coefficients of the filter coefficient matrix (H), the channel transformed coefficients being arranged to form a channel transformed matrix ($\hat{G}$).
5. The signal processing apparatus (100) of claim 4, wherein the channel determiner is configured to determine the channel transformed matrix ($\hat{G}$) according to the following equation:

$$\hat{G}(k, n) = \left(\mathbf{H}^H \mathbf{x}(k, n) \text{diag}\{X_1(k, n), X_2(k, n), \dots, X_P(k, n)\}^{-1} \right)^{-1}$$

wherein $\hat{G}$ denotes the channel transformed matrix, $\mathbf{x}$ denotes the input transformed coefficient matrix, $\mathbf{H}$ denotes the filter coefficient matrix, and X_1 to X_P denote input transformed coefficients.

6. A signal processing method (200) for dereverberating a number (Q) of input audio signals, the signal processing method (200) comprising:
 transforming (201) the number (Q) of input audio signals into a transformed domain to obtain input transformed coefficients, the input transformed coefficients being arranged to form an input transformed coefficient matrix;
 determining (203) filter coefficients upon the basis of eigenvalues of a signal space, the filter coefficients being arranged to form a filter coefficient matrix (H);
 convolving (205) input transformed coefficients of the input transformed coefficient matrix by filter coefficients of the filter coefficient matrix (H) to obtain output transformed coefficients, the output transformed coefficients being arranged to form an output transformed coefficient matrix; and
 inversely transforming (207) the output transformed coefficient matrix from the transformed domain to obtain a number of output audio signals **characterised in that:** the step of determining the filter coefficients comprises determining input auto coherence coefficients upon the basis of the input transformed coefficients, the input auto coherence coefficients indicating a coherence of the input transformed coefficients associated to a current time interval and a past time interval, the input auto coherence coefficients being arranged to form an input auto coherence matrix, and determining the filter coefficients upon the basis of the input auto coherence matrix.
7. A computer program comprising a program code for performing the signal processing method (200) of claim 6 when executed on a computer.

Patentansprüche

1. Signalverarbeitungsvorrichtung (100) zum Enthalten einer Anzahl (Q) von Eingangsaudiosignalen, wobei die Signalverarbeitungsvorrichtung (100) Folgendes umfasst:

einen Transformierer (101), ausgelegt zum Transformieren der Anzahl (Q) von Eingangsaudiosignalen in eine transformierte Domäne, um eingangstransformierte Koeffizienten zu erhalten, wobei die eingangstransformierten Koeffizienten so angeordnet werden, dass sie eine eingangstransformierte Koeffizientenmatrix bilden;
 einen Filterkoeffizienten-Bestimmer (103), ausgelegt zum Bestimmen von Filterkoeffizienten auf der Basis von Eigenwerten eines Signalraums, wobei die Filterkoeffizienten so angeordnet werden, dass sie eine Filterkoeffizientenmatrix (H) bilden;
 ein Filter (105), ausgelegt zum Falten von eingangstransformierten Koeffizienten der eingangstransformierten Koeffizientenmatrix mit Filterkoeffizienten der Filterkoeffizientenmatrix (H), um ausgangstransformierte Koeffizienten zu erhalten, wobei die ausgangstransformierten Koeffizienten so angeordnet werden, dass sie eine ausgangstransformierte Koeffizientenmatrix bilden; und
 einen Umkehr-Transformierer (107), ausgelegt zum Umkehrtransformieren der ausgangstransformierten Koeffizientenmatrix aus der transformierten Domäne, um eine Anzahl von Ausgangsaudiosignalen zu erhalten;
dadurch gekennzeichnet, dass
 der Filterkoeffizienten-Bestimmer (103) ausgelegt ist zum Bestimmen von Eingangsautokohärenzkoeffizienten auf der Basis der eingangstransformierten Koeffizienten, wobei die Eingangsautokohärenzkoeffizienten eine Kohärenz der eingangstransformierten Koeffizienten, die einem aktuellen Zeitintervall und einem vergangenen Zeitintervall zugeordnet sind, angeben, wobei die Eingangsautokohärenzkoeffizienten so angeordnet werden, dass sie eine Eingangsautokohärenzmatrix bilden, und wobei der Filterkoeffizienten-Bestimmer (103) ferner ausgelegt ist zum Bestimmen der Filterkoeffizienten auf der Basis der Eingangsautokohärenzmatrix.

2. Signalverarbeitungsvorrichtung (100) nach Anspruch 1, wobei der Filterkoeffizienten-Bestimmer (103) ausgelegt ist zum Bestimmen des Signalraums auf der Basis einer Eingangsautokorrelationsmatrix (ϕ_{xx}) der eingangstransformierten Koeffizientenmatrix.

3. Signalverarbeitungsvorrichtung (100) nach einem der vorhergehenden Ansprüche, wobei der Transformierer (101) ausgelegt ist zum Transformieren der Anzahl (Q) von Eingangsaudiosignalen in den Frequenzbereich, um die eingangstransformierten Koeffizienten zu erhalten.

4. Signalverarbeitungsvorrichtung (100) nach einem der vorhergehenden Ansprüche, ferner umfassend:
 einen Kanalbestimmer, ausgelegt zum Bestimmen von kanaltransformierten Koeffizienten auf der Basis der eingangstransformierten Koeffizienten der eingangstransformierten Koeffizientenmatrix und der Filterkoeffizienten der Filterkoeffizientenmatrix (H), wobei die kanaltransformierten Koeffizienten so angeordnet werden, dass sie eine kanaltransformierte Matrix ($\hat{G}$) bilden.

5. Signalverarbeitungsvorrichtung (100) nach Anspruch 4, wobei der Kanalbestimmer ausgelegt ist zum Bestimmen der kanaltransformierten Matrix ($\hat{G}$) gemäß der folgenden Gleichung:

$$\hat{G}(k, n) = (\mathbf{H}^H \mathbf{x}(k, n) \text{diag}\{X_1(k, n), X_2(k, n), \dots, X_P(k, n)\}^{-1})^{-1},$$

wobei $\hat{G}$ die kanaltransformierte Matrix bedeutet, $\mathbf{x}$ die eingangstransformierte Koeffizientenmatrix bedeutet, $\mathbf{H}$ die Filterkoeffizientenmatrix bedeutet und X_1 bis X_P eingangstransformierte Koeffizienten bedeuten.

6. Signalverarbeitungsverfahren (200) zum Enthalten einer Anzahl (Q) von Eingangsaudiosignalen, wobei das Signalverarbeitungsverfahren (200) Folgendes umfasst:

Transformieren (201) der Anzahl (Q) von Eingangsaudiosignalen in eine transformierte Domäne, um eingangstransformierte Koeffizienten zu erhalten, wobei die eingangstransformierten Koeffizienten so angeordnet werden, dass sie eine eingangstransformierte Koeffizientenmatrix bilden;
 Bestimmen (203) von Filterkoeffizienten auf der Basis von Eigenwerten eines Signalraums, wobei die Filterkoeffizienten so angeordnet werden, dass sie eine Filterkoeffizientenmatrix (H) bilden;
 Falten (205) von eingangstransformierten Koeffizienten der eingangstransformierten Koeffizientenmatrix mit

Filterkoeffizienten der Filterkoeffizientenmatrix (H), um ausgangstransformierte Koeffizienten zu erhalten, wobei die ausgangstransformierten Koeffizienten so angeordnet werden, dass sie ausgangstransformierte Koeffizientenmatrix bilden; und

Umkehrtransformieren (207) der ausgangstransformierten Koeffizientenmatrix aus der transformierten Domäne, um eine Anzahl von Ausgangsaudiosignalen zu erhalten,

dadurch gekennzeichnet, dass

der Schritt des Bestimmens der Filterkoeffizienten Folgendes umfasst:

Bestimmen von Eingangsautokohärenzkoeffizienten auf der Basis der eingangstransformierten Koeffizienten, wobei die Eingangsautokohärenzkoeffizienten eine Kohärenz der eingangstransformierten Koeffizienten, die einem aktuellen Zeitintervall und einem vergangenen Zeitintervall zugeordnet sind, angeben, wobei die Eingangsautokohärenzkoeffizienten so angeordnet werden, dass sie eine Eingangsautokohärenzmatrix bilden, und Bestimmen der Filterkoeffizienten auf der Basis der Eingangsautokohärenzmatrix.

7. Computerprogramm, das einen Programmcode zum Ausführen des Signalverarbeitungsverfahrens (200) nach Anspruch 6, wenn er auf einem Computer ausgeführt wird, umfasst.

Revendications

1. Appareil (100) de traitement de signaux permettant la déréverbération d'un nombre (Q) de signaux audio d'entrée, l'appareil (100) de traitement de signaux comprenant :

une unité de transformation (101) configurée pour transformer le nombre (Q) de signaux audio d'entrée dans un domaine transformé aux fins d'obtenir des coefficients transformés d'entrée, les coefficients transformés d'entrée étant agencés de manière à former une matrice de coefficients transformés d'entrée ;

une unité de détermination de coefficients de filtre (103) configurée pour déterminer des coefficients de filtre sur la base de valeurs propres d'un espace de signaux, les coefficients de filtre étant agencés de manière à former une matrice de coefficients de filtre (H) ;

un filtre (105) configuré pour convoluer des coefficients transformés d'entrée de la matrice de coefficients transformés d'entrée avec des coefficients de filtre de la matrice de coefficients de filtre (H) aux fins d'obtenir des coefficients transformés de sortie, les coefficients transformés de sortie étant agencés de manière à former une matrice de coefficients transformés de sortie ; et

une unité de transformation inverse (107) configurée pour transformer en sens inverse la matrice de coefficients transformés de sortie à partir du domaine transformé aux fins d'obtenir un nombre de signaux audio de sortie ;

caractérisé en ce que :

l'unité de détermination de coefficients de filtre (103) est configurée pour déterminer des coefficients d'autocohérence d'entrée sur la base des coefficients transformés d'entrée, les coefficients d'autocohérence d'entrée indiquant une cohérence des coefficients transformés d'entrée associés à un intervalle de temps présent et un intervalle de temps passé, les coefficients d'autocohérence d'entrée étant agencés de manière à former une matrice d'autocohérence d'entrée, et l'unité de détermination de coefficients de filtre (103) étant configurée en outre pour déterminer les coefficients de filtre sur la base de la matrice d'autocohérence d'entrée.

2. Appareil (100) de traitement de signaux selon la revendication 1, dans lequel l'unité de détermination de coefficients de filtre (103) est configurée pour déterminer l'espace de signaux sur la base d'une matrice d'autocorrélation d'entrée (Φ_{xx}) de la matrice de coefficients transformés d'entrée.

3. Appareil (100) de traitement de signaux selon l'une quelconque des revendications précédentes, dans lequel l'unité de transformation (101) est configurée pour transformer le nombre (Q) de signaux audio d'entrée dans le domaine fréquentiel aux fins d'obtenir les coefficients transformés d'entrée.

4. Appareil (100) de traitement de signaux selon l'une quelconque des revendications précédentes, comprenant en outre :

une unité de détermination de canal configurée pour déterminer des coefficients transformés de canal sur la base des coefficients transformés d'entrée de la matrice de coefficients transformés d'entrée et des coefficients de filtre de la matrice de coefficients de filtre (H), les coefficients transformés de canal étant agencés de manière à former une matrice transformée de canal ($\hat{G}$).

5. Appareil (100) de traitement de signaux selon la revendication 4, dans lequel l'unité de détermination de canal est

configurée pour déterminer la matrice transformée de canal ($\hat{G}$) selon l'équation suivante :

$$\hat{G}(k, n) = (\mathbf{H}^H \mathbf{x}(k, n) \text{diag}\{X_1(k, n), X_2(k, n), \dots, X_P(k, n)\}^{-1})^{-1}$$

où ($\hat{G}$) désigne la matrice transformée de canal, $\mathbf{x}$ désigne la matrice de coefficients transformés d'entrée, $\mathbf{H}$ désigne la matrice de coefficients de filtre, et X_1 à X_P désignent des coefficients transformés d'entrée.

6. Procédé (200) de traitement de signaux permettant la déréverbération d'un nombre (Q) de signaux audio d'entrée, le procédé (200) de traitement de signaux comprenant les étapes suivantes :

transformation (201) du nombre (Q) de signaux audio d'entrée dans un domaine transformé aux fins d'obtenir des coefficients transformés d'entrée, les coefficients transformés d'entrée étant agencés de manière à former une matrice de coefficients transformés d'entrée ;

détermination (203) de coefficients de filtre sur la base de valeurs propres d'un espace de signaux, les coefficients de filtre étant agencés de manière à former une matrice de coefficients de filtre ($\mathbf{H}$) ;

convolution (205) de coefficients transformés d'entrée de la matrice de coefficients transformés d'entrée avec des coefficients de filtre de la matrice de coefficients de filtre ($\mathbf{H}$) aux fins d'obtenir des coefficients transformés de sortie, les coefficients transformés de sortie étant agencés de manière à former une matrice de coefficients transformés de sortie ; et

transformation inverse (207) de la matrice de coefficients transformés de sortie à partir du domaine transformé aux fins d'obtenir un nombre de signaux audio de sortie ;

caractérisé en ce que :

l'étape de détermination des coefficients de filtre comprend l'étape de détermination de coefficients d'autocohérence d'entrée sur la base des coefficients transformés d'entrée, les coefficients d'autocohérence d'entrée indiquant une cohérence des coefficients transformés d'entrée associés à un intervalle de temps présent et un intervalle de temps passé, les coefficients d'autocohérence d'entrée étant agencés de manière à former une matrice d'autocohérence d'entrée, et de détermination des coefficients de filtre sur la base de la matrice d'autocohérence d'entrée.

7. Programme d'ordinateur comprenant un code de programme destiné à mettre en oeuvre le procédé (200) de traitement de signaux selon la revendication 6 lorsqu'il est exécuté sur un ordinateur.

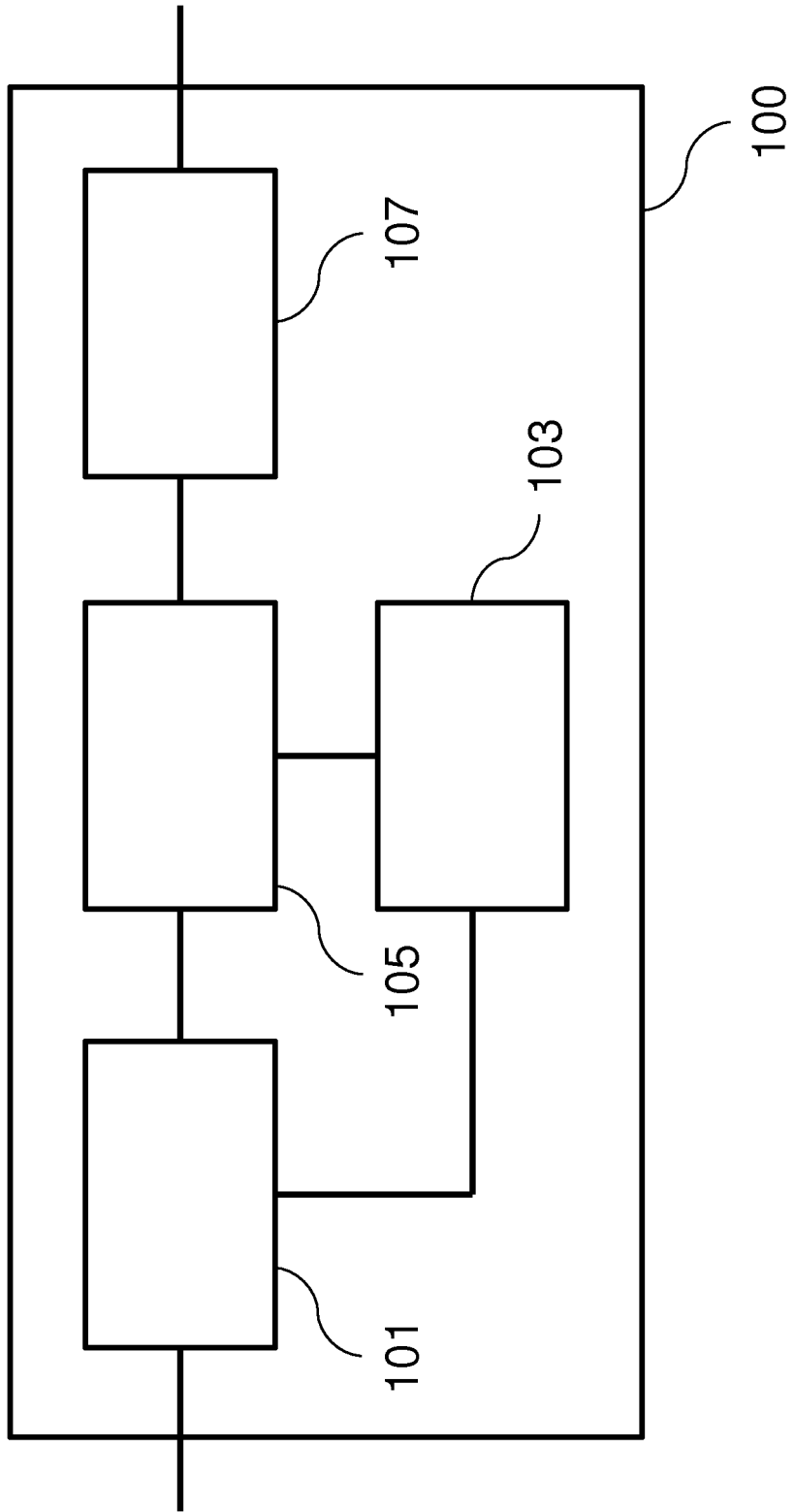


Fig. 1

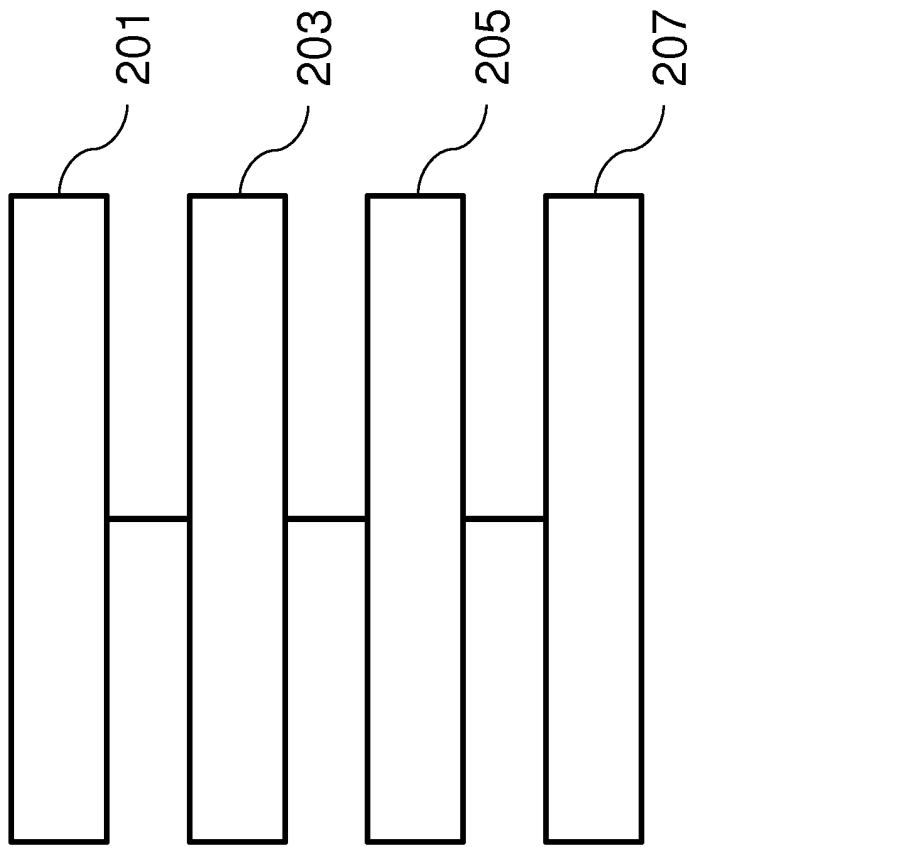


Fig. 2

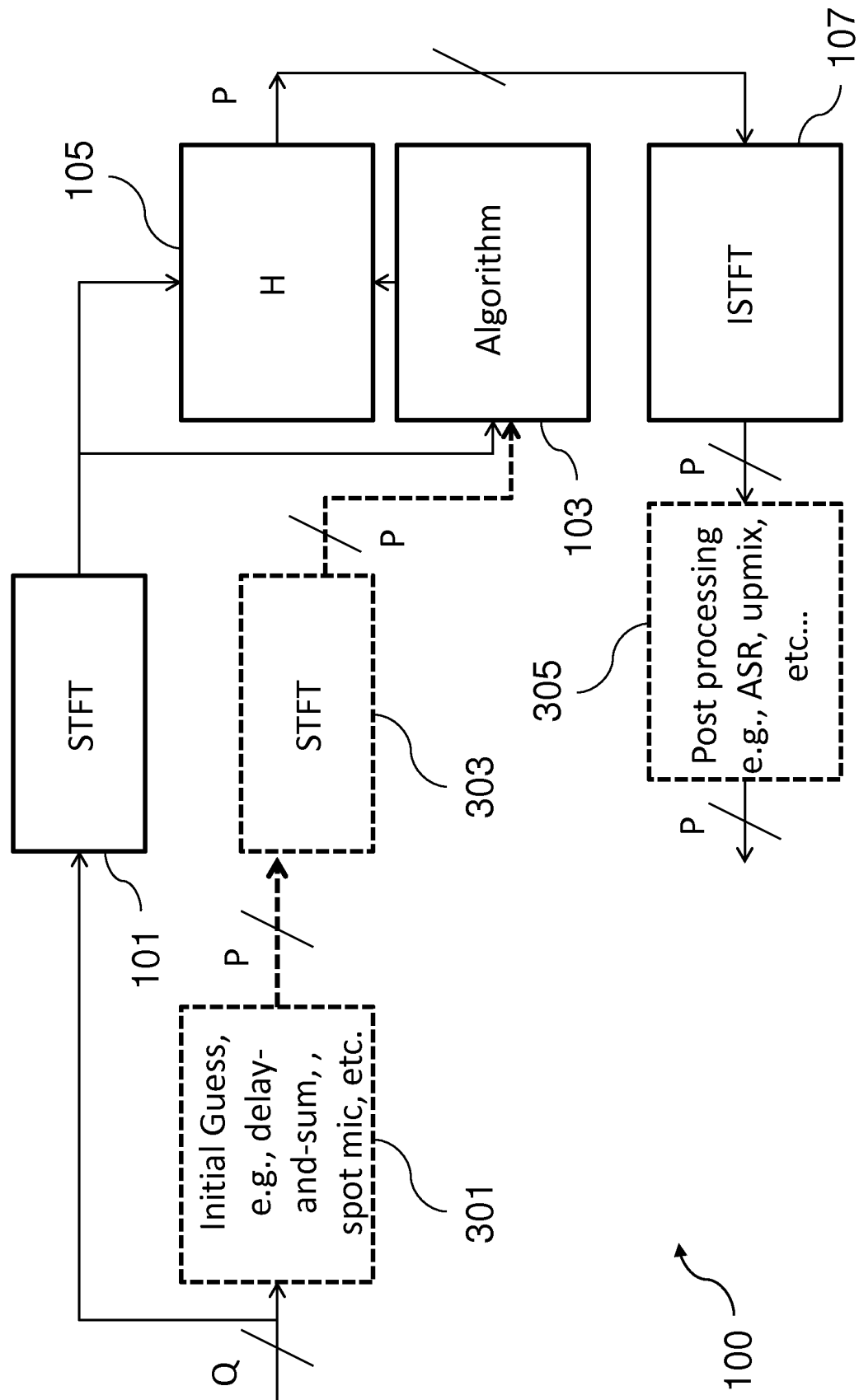


Fig. 3

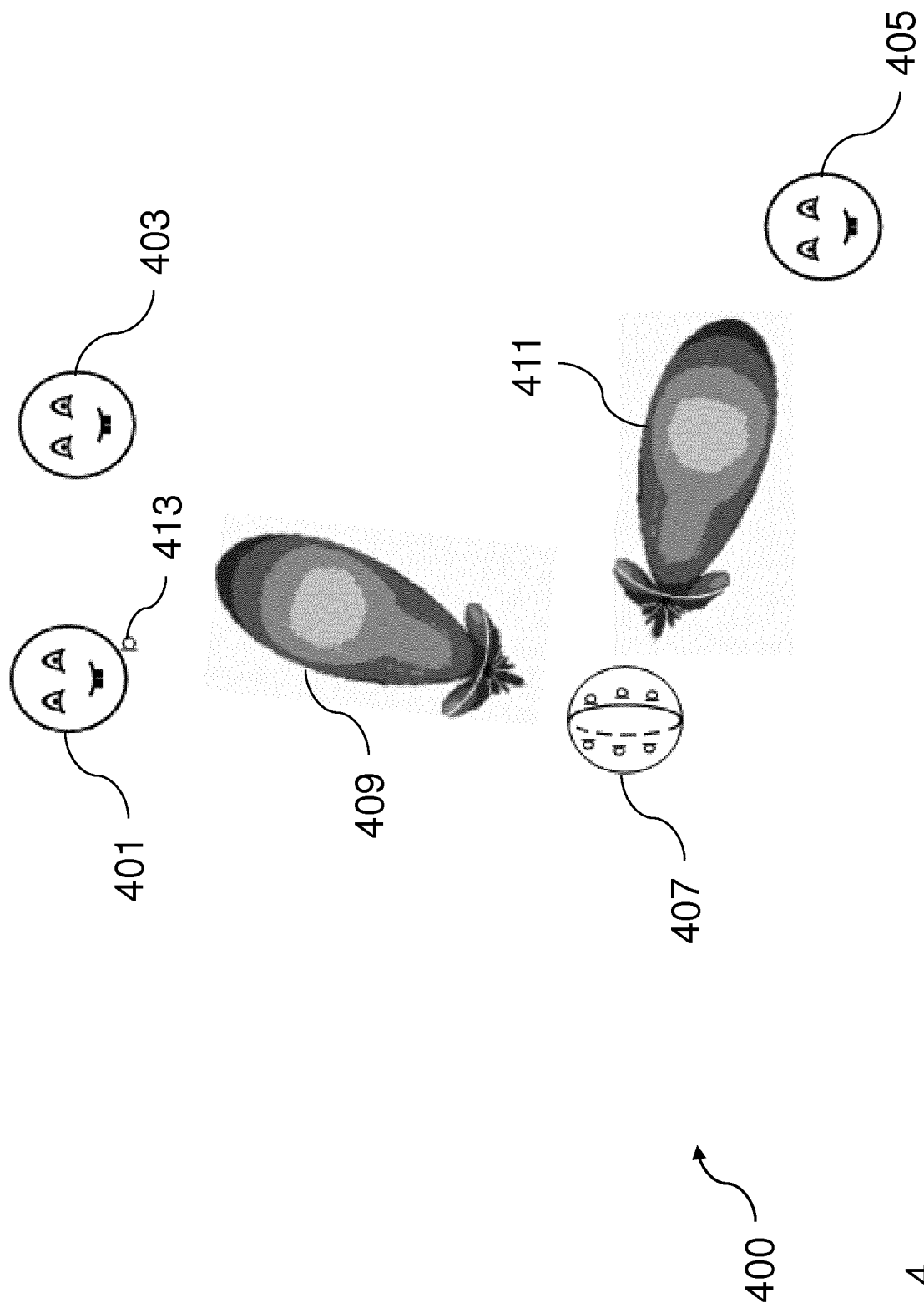


Fig. 4

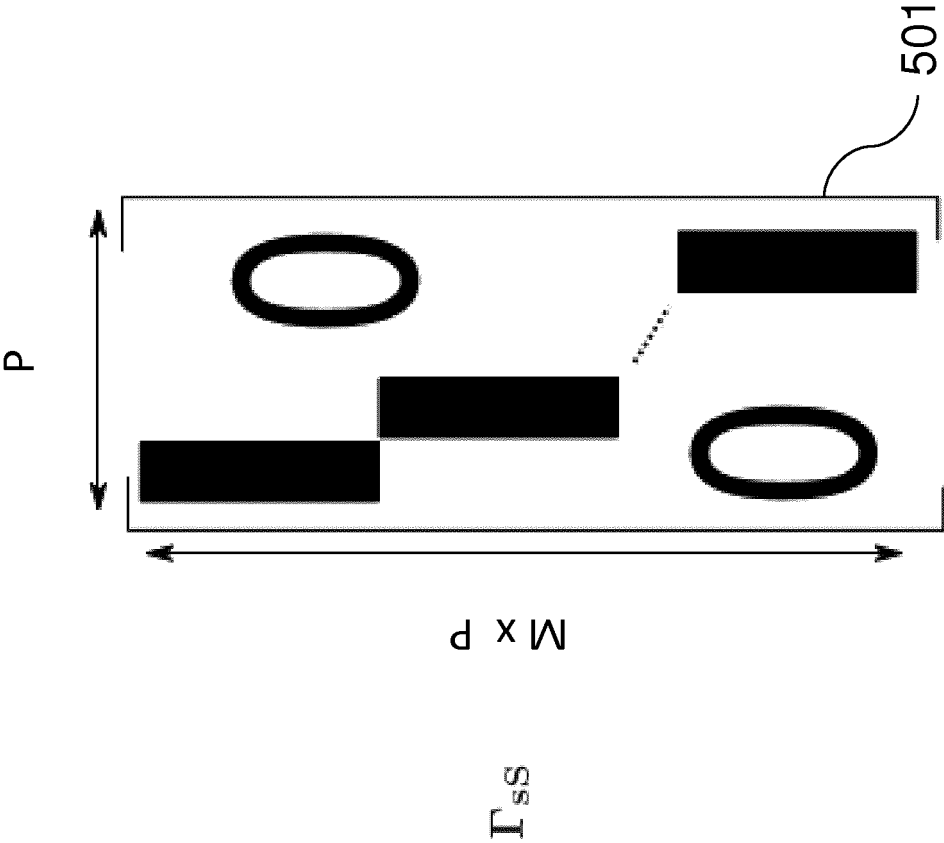


Fig. 5

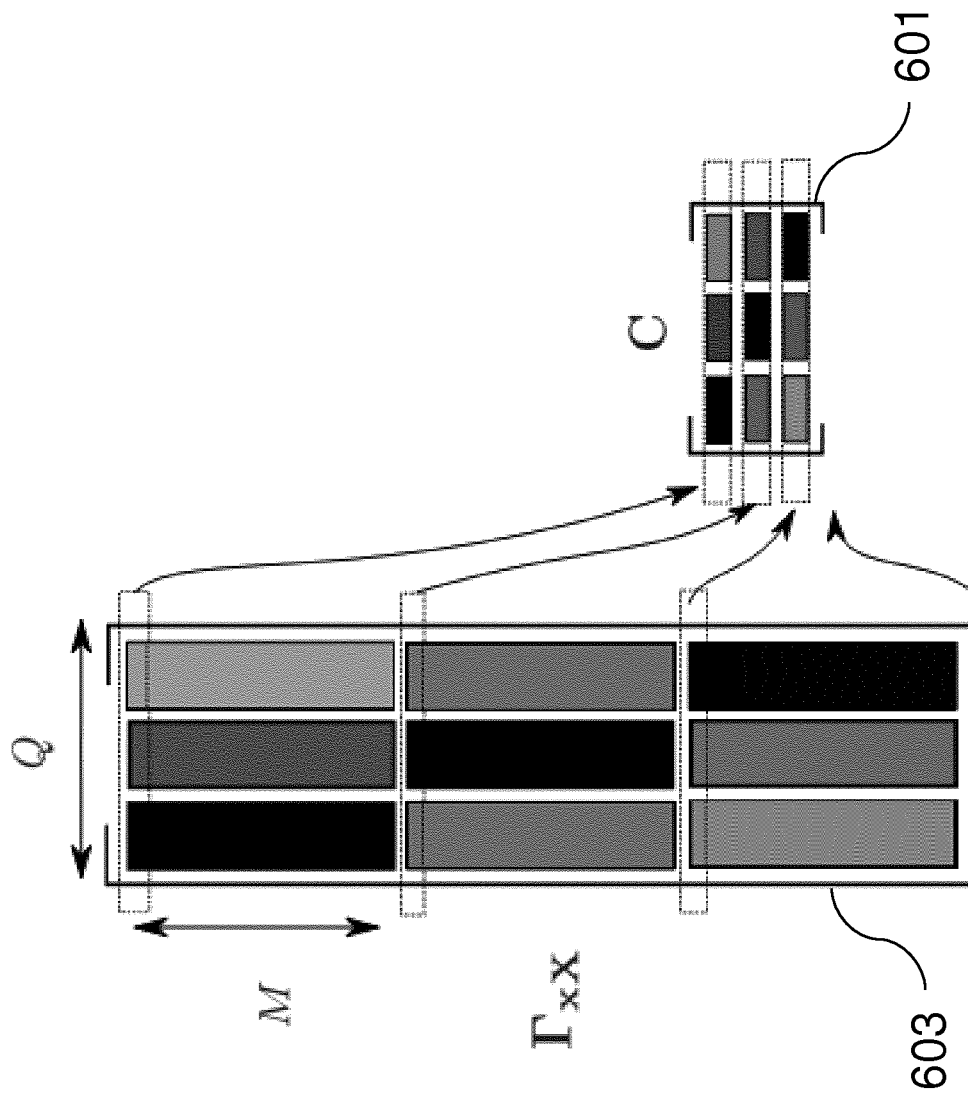


Fig. 6

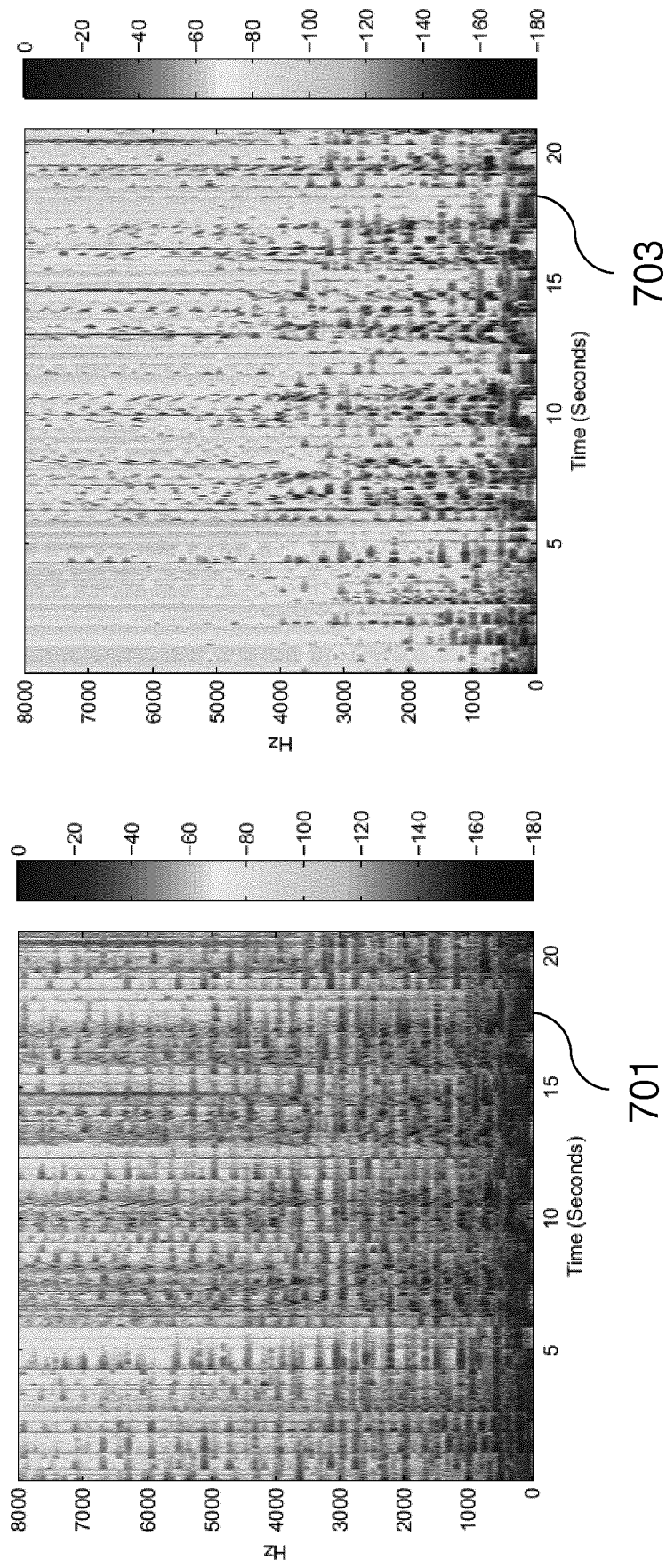


Fig. 7

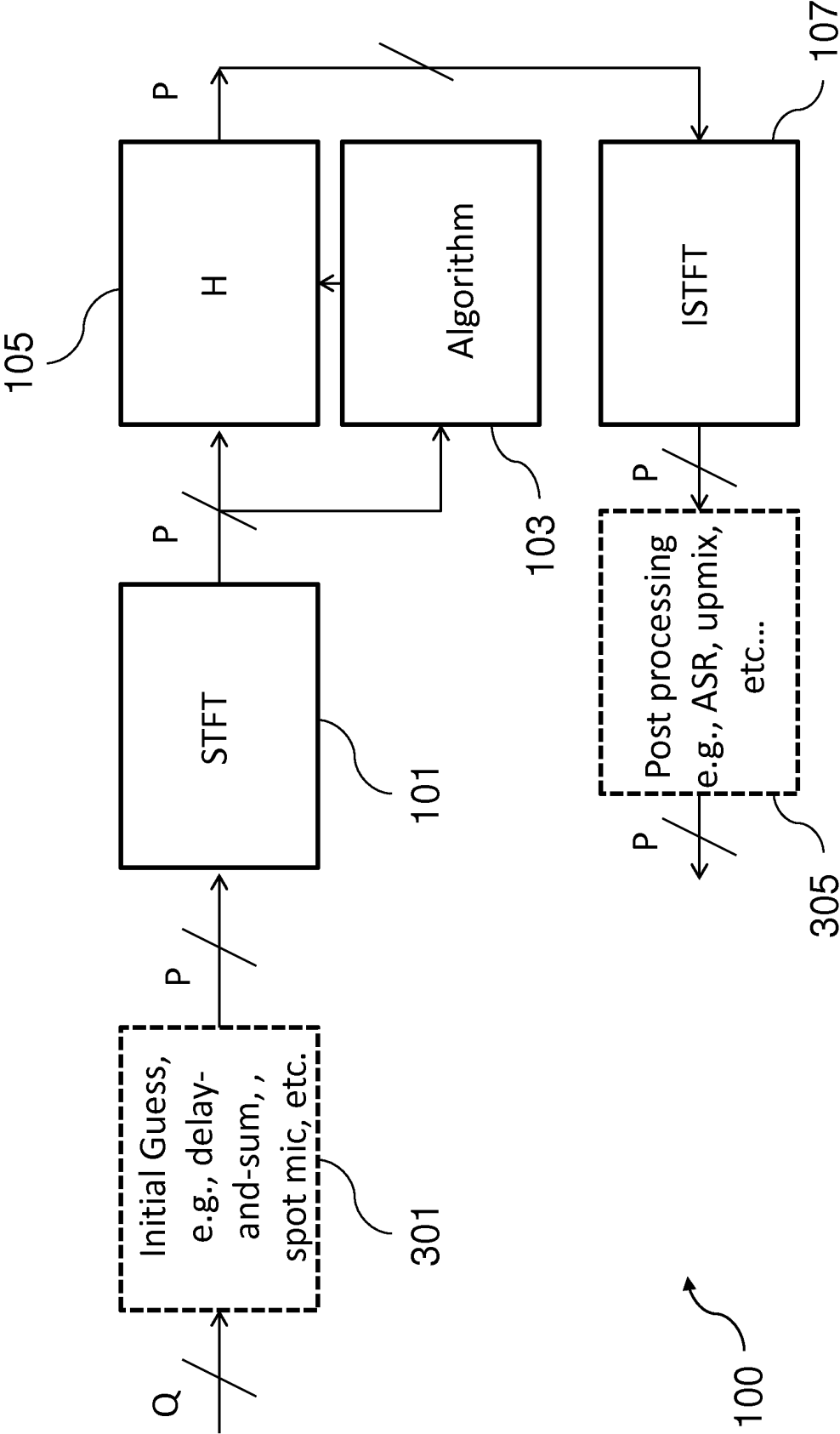


Fig. 8

REFERENCES CITED IN THE DESCRIPTION

This list of references cited by the applicant is for the reader's convenience only. It does not form part of the European patent document. Even though great care has been taken in compiling the references, errors or omissions cannot be excluded and the EPO disclaims all liability in this regard.

Non-patent literature cited in the description

- Trinicon for dereverberation of speech and audio signals. **HERBERT BUCHNER et al.** Speech Dereverberation, Signals and Communication Technology. Springer, 2010, 311-385 [0005]
- **ANDREAS WALTHER et al.** Direct-Ambient Decomposition and Upmix of Surround Signals. *IEEE Workshop on Applications of Signal Processing to Audio and Acoustics*, 2011 [0005]
- **R.S. RASHOBH ; A.W.H. KHONG ; D. LIU.** Multichannel Equalization in the KLT and Frequency Domains With Application to Speech Dereverberation. *IEEE/ACM Transactions on Audio, Speech, and Language Processing*, March 2014, vol. 22 (3 [0005]